

RESEARCH ARTICLE

OPEN ACCESS

ACO-ELM: An Ant Colony Optimization-Driven Extreme Learning Machine for Dry Bean Variety Classification

Mohit Kharbanda¹, Arshi Husain², Virendra P. Vishwakarma³

^{1,2,3}USIC&T, GGSIPU, Delhi, India

¹mohit93kharbanda@gmail.com

Abstract – Accurate classification of dry bean varieties is important for automated agricultural inspection, quality assessment, and efficient crop management. This study proposes an Ant Colony Optimization-based Extreme Learning Machine (ACO-ELM) for the classification of dry bean varieties using morphological characteristics. The Dry Bean dataset initially contained 13,611 samples with 16 numerical features representing geometric and shape-related characteristics across seven bean classes. To minimize potential data leakage, 68 duplicate records were identified and removed before data partitioning, resulting in 13,543 unique samples. The dataset was divided using a stratified 80:20 train-test split, producing 10,834 training samples and 2,709 independent test samples. An internal validation subset was further created from the training data for ACO-based hyperparameter optimization. ACO was employed to optimize the number of hidden neurons, regularization parameter, and activation function of the ELM using 20 ants over 20 iterations. The optimization selected 900 hidden neurons, a regularization parameter ($C = 50.0$), and a sigmoid activation function, achieving a validation accuracy of 94.83%. The optimized ELM was subsequently retrained using the complete training set and evaluated on the independent test set. The proposed ACO-ELM achieved 92.47% accuracy, 92.47% weighted precision, 92.47% weighted recall, and 92.46% weighted F1-score. Furthermore, the model achieved a 99.05% weighted multiclass ROC-AUC, with macro precision, recall, and F1-score of 93.65%, 93.45%, and 93.54%, respectively. The results demonstrate that ACO-based hyperparameter optimization can provide an effective ELM-based framework for multiclass dry bean variety classification.

Keywords – Dry bean classification, Extreme Learning Machine, Ant Colony Optimization, ACO-ELM, agricultural classification, hyperparameter optimization, machine learning.

1. Introduction

Agriculture plays a fundamental role in global food security, and the accurate identification and classification of crop varieties is an important component of modern agricultural management. Dry beans are among the important legume crops consumed worldwide because of their nutritional value and economic significance. Several dry bean varieties exhibit similar physical and morphological characteristics, making their manual identification difficult, time-consuming, and dependent on expert knowledge. Consequently, automated classification methods based on measurable morphological characteristics have attracted increasing attention as a means of supporting agricultural inspection, quality assessment, and variety identification.

The development of machine learning (ML) has provided effective approaches for automated classification of agricultural products. In particular, numerical morphological and geometric characteristics can provide useful information for distinguishing different dry bean varieties without requiring complex image-processing pipelines. Previous studies have investigated several conventional ML algorithms, including Support Vector Machine (SVM), Decision Tree, Random Forest, K-Nearest Neighbors (KNN), and XGBoost, for dry bean classification. Slowinski [1] evaluated several machine learning and deep learning approaches on the Dry Bean dataset and reported classification accuracies ranging from 88.35% to

93.61%. Similarly, Ardeshirifar et al. [2] investigated XGBoost and SVM for distinguishing seven dry bean varieties and demonstrated the effectiveness of machine learning models for multiclass classification. Mehta et al. [3] further examined different classification approaches and SVM kernels, with the RBF-based SVM achieving an accuracy of 93.34%. Preprocessing and class distribution management are also crucial in the classification of dry beans. According to Khan et al. [4], several multiclass classifiers were tested, and ADASYN was used for class imbalance treatment. They achieved great results using the XGBoost algorithm. Qasim et al. [5] examined the area of ensemble learning and showed that the Voting Classifier can offer better results than the separate classifiers. Therefore, one may conclude that the efficiency of the dry beans classification is highly dependent on more factors than just the classifier itself. Another significant element affecting ML models' performance is the choice of hyperparameters. Manual tuning of hyperparameters might not yield the ideal settings for the data, and exhaustive searches may be computationally costly due to increasing the search space. Naik et al. [6] conducted research on Bayesian optimization, random search, and grid search on an SVM-based dry bean classification algorithm. More generally, Nasteski [7] discussed the principles of supervised machine learning and the importance of selecting suitable learning algorithms for classification problems. Ilemobayo et al. [8] reviewed different hyperparameter opti-

mization approaches and emphasized their role in improving machine learning performance. Probst et al. [9] also demonstrated that hyperparameter importance and tunability can vary considerably across machine learning algorithms. Bergstra and Bengio [10] showed that random search can efficiently explore hyperparameter spaces compared with exhaustive grid-based approaches.

Although these studies demonstrate the effectiveness of conventional ML algorithms, ensemble methods, preprocessing techniques, and hyperparameter optimization for dry bean classification, relatively limited attention has been given to combining metaheuristic optimization with Extreme Learning Machine (ELM) for this problem. ELM is a single-hidden-layer feed-forward neural network characterized by randomly initialized input weights and analytically determined output weights. Consequently, ELM can provide rapid training while retaining the ability to model nonlinear relationships in structured numerical data. However, its performance can be influenced by important hyperparameters, particularly the number of hidden neurons, regularization parameter, and activation function. Appropriate selection of these parameters is therefore essential for obtaining a reliable ELM classifier.

Ant Colony Optimization (ACO) is a population-based metaheuristic inspired by the foraging behavior of ants. Through pheromone-based search and heuristic information, ACO can explore a predefined solution space and progressively favor promising configurations. The ability of ACO to perform adaptive search makes it suitable for hyperparameter optimization. Therefore, integrating ACO with ELM provides an opportunity to combine the fast analytical training characteristics of ELM with the adaptive search capability of a metaheuristic optimization algorithm.

Motivated by these considerations, this study proposes an **Ant Colony Optimization-based Extreme Learning Machine (ACO-ELM)** framework for multiclass classification of dry bean varieties. The proposed approach uses 16 numerical morphological features to classify seven dry bean varieties. Duplicate observations are removed before dataset partitioning to reduce the possibility of data leakage. A stratified 80:20 train-test strategy is subsequently adopted, while an internal validation subset is used exclusively for ACO-based hyperparameter optimization. ACO searches for an effective combination of the ELM hidden-layer size, regularization parameter, and activation function. The best configuration is then used to retrain the ELM using the complete training partition and is finally evaluated on an independent test set.

The main contributions of this study are summarized as follows:

- A metaheuristic-optimized ELM framework, termed ACO-ELM, is proposed for seven-class dry bean variety classification using numerical morphological characteristics.
- ACO is employed to optimize three important ELM hyperparameters, namely the number of hidden neurons,

regularization parameter, and activation function.

- Duplicate records are identified and removed before the train-test split, while the independent test set is kept completely isolated from the hyperparameter optimization process.
- The proposed framework is evaluated using accuracy, weighted precision, weighted recall, weighted F1-score, macro-averaged metrics, and multiclass ROC-AUC to provide a comprehensive assessment of classification performance.
- The proposed ACO-ELM framework achieves an independent test accuracy of 92.47% and a weighted multiclass ROC-AUC of 99.05%, demonstrating its effectiveness for automated dry bean variety classification.

The remainder of this paper is organized as follows. Section II presents the related literature on dry bean classification, machine learning, and hyperparameter optimization. Section III describes the research methodology, including dataset preprocessing, ELM formulation, ACO-based optimization, and the evaluation procedure. Section IV presents the experimental results and discussion. Section V concludes the study, while Section VI discusses potential directions for future research.

2. Literature Review

The application of machine learning (ML) for automatic classification of dry bean varieties has received increasing attention because morphological and geometric characteristics can be used to distinguish bean types efficiently. Researchers have investigated conventional classification algorithms, ensemble learning, preprocessing techniques, and hyperparameter optimization to improve classification performance. The most relevant studies are discussed below.

Slowinski [1] investigated the classification of dry bean varieties using a dataset containing more than 13,000 samples described by geometric characteristics. The study initially removed highly correlated and redundant features to obtain a more informative feature set. Several machine learning and deep learning approaches, including Multinomial Naive Bayes, Support Vector Machine (SVM), Decision Tree, Random Forest, Voting Classifier, and Artificial Neural Network, were evaluated. The reported classification accuracies ranged from 88.35% to 93.61%, demonstrating the effectiveness of geometric characteristics for automated dry bean recognition.

Ardeshtirifar et al. [2] proposed an automated dry bean classification framework based on XGBoost and SVM. Their study focused on distinguishing seven dry bean varieties having similar physical characteristics. Image-derived features were utilized, and nested cross-validation was employed for model evaluation. The authors also applied Z-score-based outlier removal, reducing the dataset to 12,909 observations, followed by dimensionality reduction to ten features. The study demonstrated

that both XGBoost and SVM can be effectively employed for multiclass dry bean classification.

Mehta et al. [3] performed a benchmarking study to examine different classification algorithms and SVM kernels for the Dry Bean Dataset. The investigation considered linear regression, polynomial regression, and SVM-based approaches. Among the evaluated configurations, the RBF kernel-based SVM obtained the highest reported accuracy of 93.34%. PCA was also applied to reduce the dimensionality of the feature space. These findings indicate that both feature reduction and appropriate kernel selection can influence the effectiveness of dry bean classification.

Khan et al. [4] compared multiple multiclass classification algorithms for the Dry Bean Dataset obtained from the UCI Machine Learning Repository. The study evaluated algorithms such as XGBoost, KNN, and Random Forest and investigated ADASYN for addressing class imbalance. XGBoost provided the strongest performance among the evaluated models, while KNN and Random Forest also demonstrated competitive results after applying the resampling technique. The study highlights the importance of combining suitable classification algorithms with appropriate preprocessing strategies.

Qasim et al. [5] investigated ensemble learning for dry bean classification and compared several classification approaches for identifying seven dry bean varieties. The Voting Classifier achieved better performance than the individual models considered in the study. The results suggest that ensemble methods can combine the predictive strengths of different classifiers and potentially provide more reliable predictions than individual machine learning models.

Naik et al. [6] investigated hyperparameter optimization for SVM-based identification of dry beans. Bayesian optimization, random search, and grid search were compared to determine suitable SVM parameter configurations. Bayesian optimization achieved a validation accuracy of 93.06% and a test accuracy of 92.29%, while random search obtained a validation accuracy of 92.62%. The study indicates that systematic optimization of model parameters can improve classification performance and provide a more objective approach to model configuration.

Nasteski [7] provided a comprehensive overview of supervised machine learning methods and explained the fundamental principles underlying supervised classification and prediction. The study discusses how learning algorithms identify relationships between input features and predefined target classes. The concepts presented in this work provide a theoretical foundation for applying supervised classification algorithms such as SVM, KNN, Decision Tree, Random Forest, and ensemble methods to structured datasets.

Hyperparameter optimization has also been recognized as an important component of machine learning model development. Illembayo et al. [8] reviewed different hyperparameter tuning approaches and emphasized their importance in improving the performance of ML algorithms. Probst et al. [9] further investigated the tunability and importance of hyperparameters

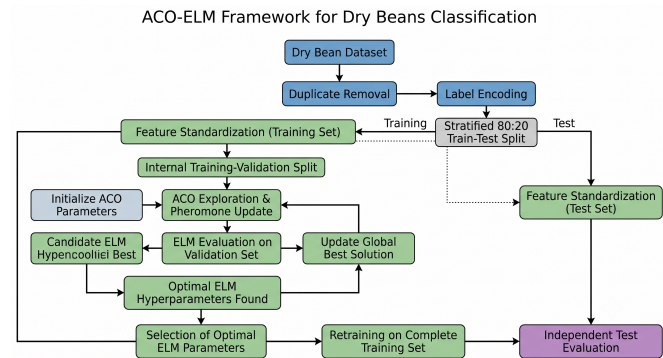


Figure 1. ACO-ELM Workflow

across different machine learning algorithms. Their findings indicate that model performance can be substantially affected by the selection of appropriate hyperparameter values, making systematic optimization preferable to arbitrary parameter selection.

Random Search was proposed by Bergstra & Bengio [10] as a method which could be used as an alternative to traditional exhaustive hyperparameter search methods. Random Search selects configurations from a distribution rather than checking each configuration of the parameters. This can be quite useful if not many hyperparameters matter a lot for the performance of the model. In general, previous research shows that classification of dry beans based on their morphology and geometry is possible using machine learning methods. Previous research has investigated classifiers, ensemble learning, feature selection, outlier elimination, class balancing, and hyperparameter tuning. However, many studies focus on a limited number of algorithms or a single optimization strategy, making it difficult to determine the most effective combination of classifier and parameter configuration. Therefore, the present study aims to provide a systematic comparative analysis of machine learning models for dry bean classification, with particular emphasis on appropriate preprocessing and hyperparameter optimization.

3. Research Methodology

The proposed methodology begins with preprocessing the raw Dry Bean dataset through duplicate removal and label encoding, followed by a stratified 80:20 train-test split and feature standardization. The training data is further split internally into training and validation sets to tune the Extreme Learning Machine (ELM) hyperparameters using Ant Colony Optimization (ACO). Finally, the optimal hyperparameters are used to retrain the ELM model on the entire training set, with ultimate generalization performance evaluated on the independent test set Figure 1.

3.1. Dataset Description

The Dry Bean dataset was employed to evaluate the proposed Ant Colony Optimization-based Extreme Learning Machine (ACO-ELM) classification framework. The orig-

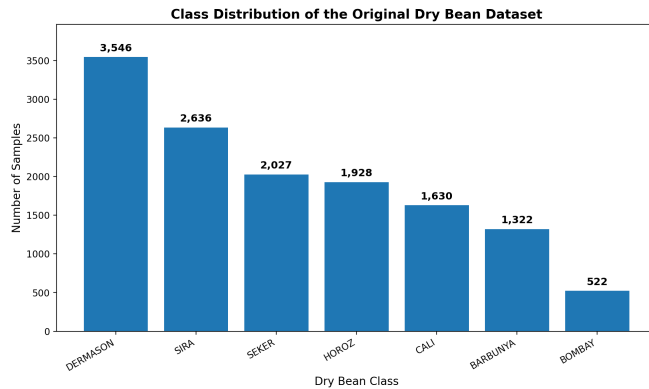


Figure 2. Class Distribution Original Dataset

inal dataset contains 13,611 samples characterized by 16 numerical morphological features Table 1 and seven bean classes, namely BARBUNYA, BOMBAY, CALI, DERMA-SON, HOROZ, SEKER, and SIRA Figure 2. The dataset contains no missing values. Before model development, 68 duplicate records were identified and removed to minimize the possibility of data leakage between the training and testing subsets. Consequently, the final dataset consisted of 13,543 unique samples.

Table 1. Class Distribution of the Dry Bean Dataset(Original Dataset

Class	Number of Samples	Percentage
DERMASON	3546	26.05%
SIRA	2636	19.37%
SEKER	2027	14.89%
HOROZ	1928	14.17%
CALI	1630	11.98%
BARBUNYA	1322	9.71%
BOMBAY	522	3.84%
Total	13611	100.00%

The 16 input features include Area, Perimeter, MajorAxisLength, MinorAxisLength, AspectRatio, Eccentricity, ConvexArea, EquivDiameter, Extent, Solidity, roundness, Compactness, ShapeFactor1, ShapeFactor2, ShapeFactor3, and ShapeFactor4. The class labels were numerically encoded using label encoding while preserving the correspondence between encoded labels and their original bean categories.

3.2. Data Preprocessing and Partitioning

After duplicate removal, the dataset was divided into training and testing subsets using a stratified 80:20 train-test split with a fixed random state of 42. Stratification was employed to preserve the relative distribution of the seven bean classes in both subsets Figure 3. This resulted in 10,834 training samples and 2,709 independent testing samples Table 2.

Feature standardization was subsequently performed using the StandardScaler transformation. The scaler was fitted exclusively on the training data and then applied to both the training

Table 2. Class Distribution of the Dry Bean Dataset(After Duplicate Removal)

Class	Samples	Percentage (%)
BARBUNYA	1,322	9.76
BOMBAY	522	3.86
CALI	1,630	12.04
DERMASON	3,546	26.18
HOROZ	1,860	13.73
SEKER	2,027	14.96
SIRA	2,636	19.46
Total	13,543	100.00

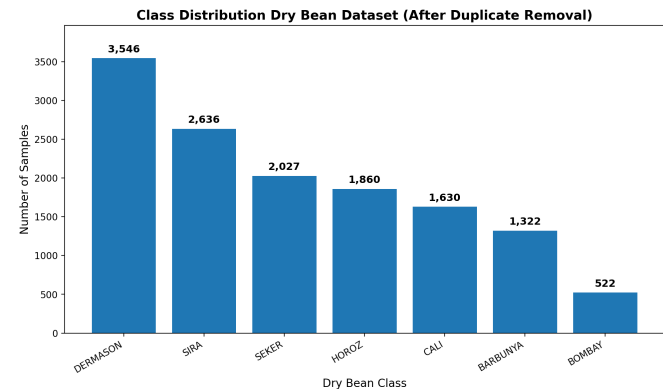


Figure 3. Class Distribution of the Dry Bean Dataset(After Duplicate Removal)

and testing data. This procedure prevents information from the independent test set from influencing the preprocessing stage.

To perform ACO-based hyperparameter optimization without accessing the independent test set, the training subset was further divided using a stratified 80:20 split. The resulting ACO-training and ACO-validation subsets contained 8,667 and 2,167 samples, respectively. The ACO-validation subset was used exclusively for evaluating candidate ELM configurations during optimization.

3.3. Extreme Learning Machine

Extreme Learning Machine (ELM) was employed as the base classifier. ELM is a single-hidden-layer feed-forward neural network in which the input weights and hidden-layer biases are randomly initialized, while the output weights are analytically determined using a regularized least-squares formulation. This eliminates the iterative gradient-based training typically required by conventional neural networks.

For an input matrix $\mathbf{X}$, the hidden-layer representation can be expressed as

$$\mathbf{H} = g(\mathbf{XW} + \mathbf{b}), \quad (1)$$

where $\mathbf{W}$ represents the randomly generated input-weight matrix, $\mathbf{b}$ denotes the hidden-layer bias vector, and $g(\cdot)$ represents the activation function. In this study, sigmoid and ReLU activation functions were considered during optimization.

The output weights were estimated using regularized least squares as

$$\beta = \left(\mathbf{H}^T \mathbf{H} + \frac{\mathbf{I}}{C} \right)^{-1} \mathbf{H}^T \mathbf{T}, \quad (2)$$

where $\mathbf{H}$ is the hidden-layer output matrix, $\mathbf{T}$ is the one-hot encoded target matrix, $\mathbf{I}$ is the identity matrix, and C is the regularization parameter. The inclusion of the regularization term improves numerical stability and controls model complexity.

3.4. Ant Colony Optimization for ELM

Ant Colony Optimization was employed to determine suitable ELM hyperparameters. In the proposed framework, ACO searches the predefined parameter space consisting of the number of hidden neurons, regularization parameter C , and activation function.

The candidate hidden-layer sizes were selected from $\{400, 500, 600, 700, 800, 900, 1000\}$ neurons. The regularization parameter C was selected from $\{0.01, 0.05, 0.1, 0.5, 1, 5, 10, 50, 100\}$. Two activation functions, sigmoid and ReLU, were considered.

Table 3. ACO Optimization Configuration

Parameter	Value
Optimization Algorithm	Ant Colony Optimization
Number of Ants	20
Number of Iterations	20
Alpha (α)	1.0
Beta (β)	2.0
Evaporation Rate	0.3
Elite Weight	2.0
Random State	42
Unique Configurations Evaluated	82
Optimization Objective	Validation Accuracy

The ACO algorithm was configured with 20 ants and 20 iterations. The parameters $\alpha = 1.0$ and $\beta = 2.0$ controlled the relative influence of pheromone information and heuristic information, respectively. A pheromone evaporation rate of 0.3 and an elite weight of 2.0 were employed Table 3. Each ant generated a candidate ELM configuration, which was trained on the ACO-training subset and evaluated on the ACO-validation subset. Validation accuracy was used as the fitness function.

For a candidate configuration s , its fitness was defined as

$$F(s) = Accuracy_{val}(s). \quad (3)$$

Following each iteration, pheromone values were reduced through evaporation and subsequently reinforced according to the quality of the candidate solutions. Elite reinforcement was additionally applied to the best solution of each iteration. This process was repeated for 20 iterations, after which the configuration having the highest validation accuracy was selected as the optimized ELM configuration.

3.5. Final ACO-ELM Training

After completion of ACO optimization, the independent test set remained completely untouched. The best hyperparameter configuration identified using the ACO-validation data was subsequently used to construct the final ELM classifier. The selected ELM was retrained using all 10,834 available training samples, thereby allowing the final model to utilize the complete training partition.

Table 4. Optimal Hyperparameters Selected by ACO

Hyperparameter	Search Space	Optimal Value
Hidden Neurons	400–1000	900
Regularization (C)	0.01–100	50.0
Activation Function	Sigmoid, ReLU	Sigmoid
Validation Accuracy	–	94.83%

The ACO process selected 900 hidden neurons, $C = 50.0$, and the sigmoid activation function, achieving a validation accuracy of 94.83% Table 4. The optimized model was then evaluated on the independent test set containing 2,709 samples.

3.6. Performance Evaluation

The performance of the final ACO-ELM model was assessed using accuracy, weighted precision, weighted recall, weighted F1-score, macro-averaged precision, macro-averaged recall, macro-averaged F1-score, and weighted multiclass ROC-AUC.

For a binary class formulation, accuracy can be expressed as

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN}, \quad (4)$$

where TP , TN , FP , and FN represent true positives, true negatives, false positives, and false negatives, respectively.

Precision, recall, and F1-score were calculated for each class and aggregated using weighted averaging to account for the different class sizes. Macro-averaged metrics were also reported to provide equal importance to each bean class.

For multiclass discrimination analysis, the weighted one-vs-rest ROC-AUC was calculated using the class probability scores generated by the ELM classifier. This provides an assessment of the model's ability to distinguish each bean class from the remaining classes independently of the final classification threshold.

The final independent test evaluation produced an accuracy of 92.47%, weighted precision of 92.47%, weighted recall of 92.47%, and weighted F1-score of 92.46%. The weighted multiclass ROC-AUC was 99.05%, while the macro precision, macro recall, and macro F1-score were 93.65%, 93.45%, and 93.54%, respectively.

4. Experimental Results and Discussion

4.1. Experimental Results

The proposed Ant Colony Optimization-based Extreme Learning Machine (ACO-ELM) was evaluated on the Dry Bean

dataset after removing 68 duplicate records. The resulting dataset contained 13,543 unique samples and 16 numerical features distributed across seven bean classes. A stratified 80:20 train-test split produced 10,834 training samples and 2,709 independent test samples. The training subset was further divided into 8,667 ACO-training samples and 2,167 validation samples for hyperparameter optimization.

ACO was configured with 20 ants and 20 iterations. The search process optimized three ELM parameters: the number of hidden neurons, the regularization parameter C , and the activation function. The best configuration obtained during optimization consisted of 900 hidden neurons, $C = 50.0$, and the sigmoid activation function. This configuration achieved a validation accuracy of 94.83%.

After optimization, the selected ELM configuration was re-trained using all 10,834 training samples and evaluated on the independent test set. The final performance is presented in . The results demonstrate consistent performance across the

Table 5. Class-Wise Performance of the Proposed ACO-ELM Model

Class	Precision (%)	Recall (%)	F1-Score (%)	Support
BARBUNYA	93.70	89.81	91.71	265
BOMBAY	100.00	100.00	100.00	104
CALI	92.84	95.40	94.10	326
DERMASON	91.62	92.52	92.07	709
HOROZ	95.37	94.09	94.72	372
SEKER	93.92	95.07	94.49	406
SIRA	88.12	87.29	87.70	527
Macro Average	93.65	93.45	93.54	2709
Weighted Average	92.47	92.47	92.46	2709

principal evaluation metrics. The difference between weighted precision, recall, and F1-score is very small, indicating that the classifier maintains relatively balanced predictive performance across the test samples. The macro-averaged F1-score of 93.54% further indicates that the model performs effectively across the individual classes rather than relying predominantly on the larger classes.

4.2. Class-Wise Performance Analysis

The detailed class-wise performance of ACO-ELM is presented in Table 5.

The model achieved its highest performance for the BOMBAY class, obtaining 100% precision, recall, and F1-score. However, BOMBAY represents the smallest test class with 104 samples; therefore, the perfect score should be interpreted in the context of its relatively limited sample size.

Among the remaining classes, HOROZ achieved an F1-score of 94.72%, followed by SEKER at 94.49% and CALI at 94.10%. DERMASON achieved an F1-score of 92.07%, while BARBUNYA achieved 91.71%. The comparatively lowest performance was obtained for SIRA, with precision, recall, and F1-score of 88.12%, 87.29%, and 87.70%, respectively.

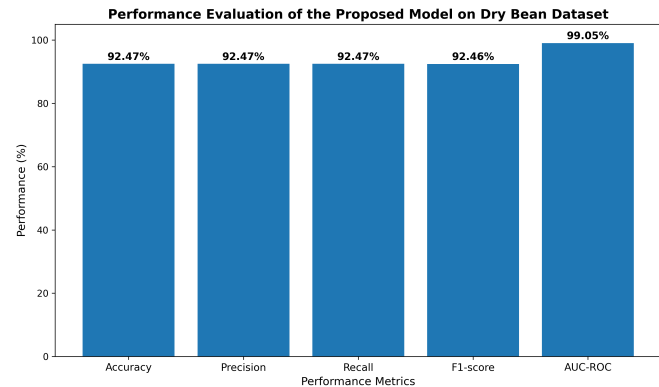


Figure 4. Performance of Proposed Model ACO-ELM

The lower performance for SIRA can be further explained using the confusion matrix. A considerable number of SIRA samples were incorrectly classified as DERMASON, indicating that these two classes exhibit relatively similar morphological characteristics in the feature space.

4.3. Discussion

The experimental findings indicate that the proposed ACO-ELM framework provides effective multiclass classification of dry bean varieties using numerical morphological characteristics. The independent test accuracy of 92.47% demonstrates that the optimized ELM can successfully discriminate among seven bean classes without using image-based deep learning techniques.

The optimization process selected a relatively large hidden layer of 900 neurons together with $C = 50.0$ and sigmoid activation. The repeated selection of this configuration during later ACO iterations indicates that the pheromone mechanism converged toward a stable region of the search space. The sigmoid activation may have been beneficial because it provides a smooth nonlinear transformation of the standardized morphological features, allowing the ELM to model nonlinear relationships among bean varieties.

The difference between the ACO validation accuracy (94.83%) and independent test accuracy (92.47%) is approximately 2.36 percentage points. Such a difference indicates some variation between the internal validation and independent test subsets, but the test performance remains strong. Importantly, the independent test set was isolated before ACO optimization and was not used for selecting the ELM parameters. Therefore, the reported test results provide a more reliable estimate of the model's generalization performance.

The macro F1-score of 93.54% is higher than the weighted F1-score of 92.46%, suggesting that the model does not depend exclusively on the majority classes. The class-wise results also demonstrate relatively consistent performance across the seven categories, although SIRA remains the most challenging class. The confusion between SIRA and DERMASON suggests that additional feature engineering or more discrim-

inative morphological descriptors could potentially improve classification performance.

The weighted multiclass ROC-AUC of 99.05% further indicates strong class-discrimination capability based on the model's class probability scores. Nevertheless, ROC-AUC and accuracy measure different aspects of classifier performance; therefore, the high ROC-AUC should not be interpreted as equivalent to 99.05% classification accuracy.

Overall, the results demonstrate that combining Ant Colony Optimization with ELM provides an effective framework for selecting ELM hyperparameters and achieving reliable dry bean variety classification. The use of duplicate removal before dataset partitioning and an untouched independent test set also strengthens the reliability of the experimental evaluation.

5. Conclusion

This study presented an Ant Colony Optimization-based Extreme Learning Machine (ACO-ELM) for multiclass classification of dry bean varieties using morphological characteristics. The Dry Bean dataset initially contained 13,611 samples and 16 numerical features representing geometric and shape-related properties of seven bean classes. To reduce the possibility of data leakage, 68 duplicate records were removed before data partitioning, resulting in 13,543 unique samples. A stratified 80:20 train-test strategy was adopted, while an internal validation subset of the training data was used exclusively for ACO-based hyperparameter optimization.

The proposed ACO framework optimized the number of hidden neurons, regularization parameter, and activation function of the ELM. The optimization process identified 900 hidden neurons, $C = 50.0$, and sigmoid activation as the best configuration, achieving a validation accuracy of 94.83%. After optimization, the ELM was retrained using the complete training set and evaluated on an independent test set. The proposed ACO-ELM achieved 92.47% accuracy, 92.47% weighted precision, 92.47% weighted recall, and 92.46% weighted F1-score. In addition, the model obtained a 99.05% weighted multiclass ROC-AUC, demonstrating strong discriminative capability.

The class-wise analysis showed strong performance for most bean varieties, with BOMBAY achieving 100% precision, recall, and F1-score. The problem was mainly encountered when classifying SIRA and DERMASON because of the similarity of their morphologies. Regardless of the problem, the high macro F1-score of 93.54% shows that the model still performed well on the seven classes. On the whole, the experimental results indicate that the ACO algorithm is efficient enough to tune hyperparameters of ELM and to offer a robust machine learning platform for classifying dry beans. It must be noted that the use of a nonlinear feature representation with the help of ELM combined with population-based hyperparameter tuning with the help of ACO provides an efficient solution as compared to more sophisticated classifiers. The future research can explore cross-validation-based tuning, feature selection algorithms, ensemble learning, among others.

6. Future Work

Despite the good results obtained by the suggested ACO-ELM approach for the classification of dry beans, there are many areas that can be examined to increase the efficiency of the suggested classifier. For instance, in the future work, the researchers could utilize stratified k -fold cross-validation in the assessment of ACO's fitness, rather than use the single validation set. Secondly, the present model seeks to optimize the number of hidden neurons, regularization parameter and the activation function. Future research on this model may seek to expand the parameter search space for other ELM parameters such as the method of initializing input weights. Feature-selection mechanisms can also be incorporated into ACO to identify the most discriminative morphological characteristics while simultaneously optimizing the ELM architecture.

Third, the proposed ACO-ELM can be compared systematically with other metaheuristic-optimized ELM variants, such as PSO-ELM, GA-ELM, Grey Wolf Optimizer-based ELM, and other swarm-intelligence approaches. Such comparisons would provide a clearer assessment of the contribution of ACO to ELM hyperparameter optimization.

Fourth, the relatively lower performance observed for the SIRA class suggests that future research should investigate class-specific feature engineering and ensemble strategies to reduce confusion between morphologically similar classes, particularly SIRA and DERMASON. Additional statistical and morphological descriptors may provide more discriminative information for these classes.

Finally, future investigations can evaluate the proposed framework on larger and more diverse agricultural datasets collected under different environmental and acquisition conditions. Integration of the optimized ELM with image-based features and deep feature extraction could also be explored to develop a hybrid classification framework capable of combining conventional morphological descriptors with learned visual representations. These extensions may improve the robustness, scalability, and practical applicability of the proposed ACO-ELM approach for automated agricultural variety classification.

Conflict of Interest

The authors declare no competing interests related to this work.

Acknowledgment

This research was supported by a fellowship under the Visvesvaraya PhD Scheme, which the authors gratefully acknowledge.

References

- [1] G. Slowinski, "Dry Beans Classification Using Machine Learning," *29th International Workshop on Concurrency, Specification and Programming (CS&P'21)*, CEUR Workshop Proceedings, 2021.

- [2] R. Ardeshirifar, “Automated Classification of Dry Bean Varieties Using XGBoost and SVM Models,” arXiv:2408.01244v1 [cs.LG], Aug. 2, 2024.
- [3] A. Mehta, P. Sengupta, D. Garg, H. Singh, and Y. S. Diamand, “Benchmarking the Effectiveness of Classification Algorithms and SVM Kernels for Dry Beans,” arXiv:2307.07863, 2023.
- [4] M. S. Khan, T. D. Nath, M. S. Hossain, A. Mukherjee, H. B. Hasnath, T. M. Meem, and U. Khan, “Comparison of Multiclass Classification Techniques Using Dry Bean Dataset,” *International Journal of Cognitive Computing in Engineering*, 2023.
- [5] Q. Zaheer, A. U. Rehman, S. Javed, T. M. Ali, and A. Mir, “Leveraging Ensemble Learning for Dry Beans Classification,” *International Conference on Engineering & Computing Technologies*, May 2024.
- [6] N. K. Naik, P. K. Sethy, R. Amat, S. K. Behera, and P. Biswas, “Evaluation of Optimization Techniques with Support Vector Machine for Identification of Dry Beans,” *Institute of Advanced Engineering and Science (IAES)*, vol. 32, no. 2, 2023.
- [7] V. Nasteski, “An Overview of the Supervised Machine Learning Methods,” Dec. 2017.
- [8] J. A. Ilemobayo, O. I. Durodola, T. O. Adewumi, O. Falana, et al., “Hyperparameter Tuning in Machine Learning: A Comprehensive Review,” *Journal of Engineering Research and Reports*, vol. 26, no. 6, pp. 388–395, June 2024.
- [9] P. Probst, A.-L. Boulesteix, and B. Bischl, “Tunability: Importance of Hyperparameters of Machine Learning Algorithms,” *Journal of Machine Learning Research*, vol. 20, pp. 1–32, 2019.
- [10] J. Bergstra and Y. Bengio, “Random Search for Hyperparameter Optimization,” *Journal of Machine Learning Research*, vol. 13, pp. 281–305, 2012.